
NetGent: Agent-Based Automation of Network Application Workflows

Jaber Daneshamooz* Eugene Vuong[‡] Laasya Koduru* Sanjay Chandrasekaran* Arpit Gupta*

*University of California Santa Barbara [‡]California State University, East Bay

Abstract

We present *NetGent*, an AI-agent framework for automating complex application workflows to generate realistic network traffic datasets. Developing generalizable ML models for networking requires data collection from network environments with traffic that results from a diverse set of real-world web applications. However, using existing browser automation tools that are diverse, repeatable, realistic, and efficient remains fragile and costly. *NetGent* addresses this challenge by allowing users to specify workflows as natural-language rules that define state-dependent actions. These abstract specifications are compiled into nondeterministic finite automata (NFAs), which a state synthesis component translates into reusable, executable code. This design enables deterministic replay, reduces redundant LLM calls through state caching, and adapts quickly when application interfaces change. In experiments, *NetGent* automated more than 50+ workflows spanning video-on-demand streaming, live video streaming, video conferencing, social media, and web scraping, producing realistic traffic traces while remaining robust to UI variability. By combining the flexibility of language-based agents with the reliability of compiled execution, *NetGent* provides a scalable foundation for generating the diverse, repeatable datasets needed to advance ML in networking.

1 Introduction

Machine learning for networking has become an increasingly active area of research for tasks such as QoE inference and optimization of networked applications. A persistent barrier is access to realistic, labeled application data at scale [1]. Unlike vision or NLP, networking datasets often *cannot* be scraped or passively collected: they must be *generated* by executing live application workflows (e.g., streaming a video, joining a meeting, browsing social media) so that traffic, logs, and user interactions reflect real deployments. Today, researchers commonly rely on browser-automation scripts (e.g., Selenium [22], PyAutoGUI [32]) to repeat experiments and scale data collection. However, authoring and maintaining such scripts is long, manual, and brittle, especially for complex, multi-step tasks across diverse sites.

As a result, prior work frequently narrows the scope to a small set of applications or use cases when generating datasets [4]. Yet building generalizable ML models requires collecting across *many* applications, inputs (e.g., different videos), and network environments [7]. To keep up with this demand, data collection pipelines must repeatedly execute the same workflows under varied conditions while remaining robust to evolving user interfaces and behavior.

Consider a concrete example: automate Disney+ to open ESPN, select the first video, and move the playback slider to the five-minute mark. Even this simple task varies: a user may or may not be logged in; profile selection may be required; the ESPN entry point and page layout change over time; ads or PIN prompts may appear. Such variability makes it difficult to design automation that is both robust and repeatable, especially when scaled to thousands or millions of runs across diverse network conditions for ML training and evaluation.

Requirements. This example illustrates six interrelated requirements for networking data generation: (1) *diversity* across applications and platforms; (2) *repeatability* so identical inputs yield identical outcomes across many runs and network conditions; (3) *complexity* to capture dynamic, non-linear, multi-step interactions; (4) *robustness* to survive frequent UI changes; (5) *realism* to mimic human behavior and avoid undesired bot detection; and (6) *efficiency* to minimize token usage and workflow-generation time. Meeting all six simultaneously is non-trivial; improving one dimension often degrades another.

Why existing approaches fall short. Web/GUI agents and script-based automation each solve a subset of these needs. Agentic approaches (e.g., ReAct [42], Reflexion [30]) emphasize online planning and self-reflection, but incur high token costs per execution and remain unreliable on long horizons—GPT-4 agents achieve $\approx 14\%$ success on WebArena versus $\approx 78\%$ for humans [43]; Mind2Web [8], VisualWebArena [16], and BrowserGym [5] further document these gaps. Scripted frameworks (Selenium, Playwright, PyAutoGUI) provide efficient replay but are notoriously flaky under UI drift; empirical studies attribute failures to asynchronous waits, DOM instability, and timing issues [19, 35, 24]. More specialized web scraping tools, such as for the broadband-plan querying tool (BQT [24]), trade generality for robustness, but are still fragile to UI changes. None of these simultaneously delivers diversity, repeatability at scale, robustness, realism, and efficiency.

Proposed approach. We introduce *NetGent*, an AI-agent framework that separates *what* a workflow should do from *how* it is executed. Users provide natural-language state prompts—high-level *trigger-action* rules (e.g., “if on login page, enter credentials,” “if viewing profiles, select `sn1client`,” “if cookies popup, select Accept”)—which specify an abstract, non-linear workflow. A *State Synthesis* component compiles these abstract prompts into *concrete states* with application-bound detectors (DOM/text/URL) and reusable executable code. Concrete states are cached in a repository and deterministically replayed by a *State Executor*; when UIs change, *NetGent* regenerates only the affected states from the same abstract prompts. This compile-then-replay design blends the flexibility of language-based synthesis with the efficiency and stability of compiled execution, directly addressing the six requirements above.

Contributions and evidence. We implement *NetGent* and evaluate it across 50+ workflows spanning video-on-demand streaming, live video streaming, video conferencing, social media, and web scraping (similar to BQT). Section 2 details the abstractions and execution model; Section 3 demonstrates that (i) abstract user prompts in natural language expand into hundreds of lines of executable code across diverse applications (extensibility), (ii) caching and replay reduce token cost and make millions of repeat runs economically feasible (efficiency and repeatability), and (iii) UI drift is handled by regenerating only the impacted state (robustness). Together, these results position *NetGent* as a proof of concept for scalable and realistic data generation in networking, complementing controllable platforms such as netReplica [7]—capable of emulating a diverse range of realistic networking conditions.

2 System Design

2.1 Architectural Abstractions

NetGent separates *what* a workflow should do from *how* it is executed through three abstractions.

Abstract NFA. Users define an abstract nondeterministic finite automaton (NFA) [25] using natural-language state prompts. Each state prompt specifies *triggers* (conditions that identify the state), *actions* (intended task), and an optional *end condition*. This representation captures non-linear flows (complexity) while keeping intent decoupled from UI specifics (robustness). For example, in a Disney+/ESPN workflow, states may include `login`, `select_profile`, `navigate_to_espn`, `select_video`, and `playback`.

Concrete NFA. During execution, *NetGent* compiles each abstract state into a *concrete state* defined by $\hat{s} = (\text{detectors}, \text{code})$: a set of CSS element, text, or URL detectors bound to the current application version, together with reusable executable code. This compiled form enables deterministic replay (repeatability) and cross-run reuse (efficiency). For example, the abstract trigger “if on login page” becomes a detector set (form labels, button text, stable DOM paths) and a short program that types credentials and clicks “Log In.”

Cache and Replay. Concrete states are stored in a *State Repository*; a *State Executor* replays their code deterministically. If a detector later fails due to UI drift, only that state is regenerated from the abstract rule (robustness). Common states (e.g., login, select_profile) are reusable across workflows and apps (efficiency, diversity).

2.2 Workflow Execution Model

Figure 1 illustrates the runtime loop which generates executable code from user prompts.

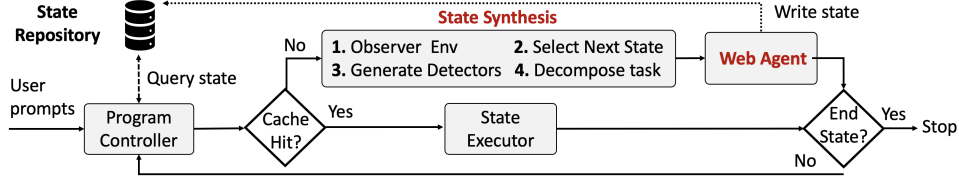


Figure 1: *NetGent* runtime loop progresses from the initial to the end state

Controller queries the cache first. Given the current page (DOM) and the last transition, the *Program Controller* queries the *State Repository*. If a cache hit occurs, the Controller invokes the *State Executor* to replay the stored code in the browser and the workflow advances. This cache-first policy is the core of compile-then-replay and eliminates repeated reasoning (repeatability, efficiency).

Cache miss triggers one-shot synthesis. On a cache miss, the Controller invokes *State Synthesis* (LLM), which performs four steps using the current DOM, screenshot, and user’s rules: (1) *Observe* the environment to form a structured view; (2) *Select* the appropriate next abstract state (trigger-action pair); (3) *Generate* concrete detectors that reliably recognize that state; (4) *Decompose* the action into a simplified plan with decomposed tasks. The decomposed tasks are then executed by the *Web Agent* which also generate the executable code. Then, the *Concrete State* is written back to the repository. Only the missing node is synthesized; the abstract NFA and prior states remain intact (robustness, efficiency).

Realistic execution and termination. To enhance realism and evade bot detection, our web agent integrates browser stealth, human-like interaction, and network control (details in Appendix §A.2). An end state is declared when an application-level condition holds (e.g., for ESPN, a <video> element is playing and time is advancing). Otherwise, the Controller loops to the next state.

Concrete example. Starting at the Disney+ homepage, the Controller hits cached login and select_profile states on subsequent runs; on the first run these are synthesized once. Navigating to the ESPN hub and clicking the first video may trigger ads or a PIN prompt; the NFA branches handle these cases by synthesizing (once) a type_pin or skip_ad state and writing them to the repository. Playback detection serves as the end state, after which *NetGent* records the successful trace and terminates.

3 Evaluation

We evaluate *NetGent* against the requirements introduced in §2, mapping experiments to the abstractions in §2.1 using Gemini 2.5 [11] with a temperature of 0.2. Details of the APIs and frameworks employed are provided in §A.1. All evaluations and executions were conducted on a MacBook Pro (Apple M3 Pro chip, 11-core CPU, 14-core GPU, and 18 GB RAM).

Diversity across applications. A central goal of *NetGent* is to keep user effort low while scaling to diverse applications. We hypothesize that prompt length serves as a proxy for user effort, while the lines of generated code reflect the automation complexity that would otherwise need to be implemented manually. We therefore measured the size of user prompts and the length of generated Python code across 50+ workflows spanning four domains: video-on-demand streaming, live video streaming, video conferencing, social media, and web scraping. Table 1 in the appendix provides the list of these applications, along with their evaluation based on code generation time, number of tokens used, and dollar cost. Each workflow requires only a 100–200 word prompt to generate code that spans hundreds of lines. This large expansion factor demonstrates that small, uniform specifications

suffice to produce substantial executable workflows. Moreover, the same prompt structure generalizes across platforms within a category (e.g., Hulu and Disney+), showing that *NetGent* is easily extensible to new applications with minimal effort.

Efficiency and repeatability. Having established diversity, we next examine efficiency and repeatability within a specific workflow. The system leverages the compile–then–replay method to achieve repeatability by reusing stored concrete states, while caching reduces token usage by avoiding redundant LLM calls.

We focus on the ESPN workflow, where user interactions include `login`, `select_profile`, `playback` and other related actions. Running this workflow without any stored concrete NFA consumes 278k tokens per run, translating to \$0.098 at \$0.35 per million tokens¹. Executing this workflow one million times without reusing a concrete NFA would cost roughly \$98,000 in LLM usage alone. By contrast, the compile–then–replay approach reuses stored states, eliminating per-run LLM costs and ensuring deterministic execution. This avoids redundant LLM and API calls, allowing a single generated workflow to be reused across multiple runs.

However, due to UI changes and other factors, periodic updates are necessary to handle new or modified states. Assuming a system with 10 states, generating each new state requires ≈ 42.6 k tokens ($\approx \$0.015$) on average. Generating workflows for all 10 states from scratch incurs a one-time cost of $\approx \$0.15$. If the workflow drifts weekly over a year (52 weeks), updating all states would cost $\approx \$7.8$. With caching, only changed states need updating. Assuming one state update per week, the annual LLM cost drops to $\approx \$0.78$. Caching thus minimizes redundant LLM calls and enables deterministic replay of previously synthesized states. This demonstrates that compile–then–replay with state caching ensures deterministic behavior while making large-scale execution economically viable.

Robustness under UI drift. Finally, we evaluate robustness to interface changes. The hypothesis is that *NetGent* can localize regeneration to only the affected state, avoiding costly re-synthesis of the full workflow. We perturb the ESPN workflow, requiring a PIN for profile access. In this case, only the affected state (`type_pin`) was regenerated, while other states such as `login`, `select_profile` and `navigate_to_espn` were replayed from cache without modification. Using the caching method, regenerating the affected state required only ≈ 20 k tokens and the entire process took 216 seconds, compared to ≈ 375 k tokens and 406 seconds if done from scratch. This bounded overhead confirms that state-level regeneration is sufficient. Even when multiple states change, the unaffected portions of the workflow remain intact. Such robustness ensures that workflows remain usable despite frequent UI drift in production applications. See Figure 2 in the appendix for the full ESPN workflow.

Summary. Across all experiments, *NetGent* satisfies the requirements of §2: prompts abstract away application-specific details to ensure diversity and extensibility; caching enables efficiency and repeatability; and state-local regeneration ensures robustness to UI drift. These results collectively demonstrate that *NetGent* provides a scalable foundation for generating realistic networking datasets across heterogeneous applications.

4 Limitations and Future Work

While *NetGent* demonstrates that abstract NFAs combined with compile–then–replay enable scalable and repeatable workflows, several limitations remain. First, manual workflow verification and failure handling currently require user intervention; automating step-level validation and state-level recovery would enable self-healing and fully autonomous workflows. Second, *NetGent* is limited to web applications; extending the NFA abstraction to desktop environments would broaden applicability. Together, these extensions will make *NetGent* more autonomous, robust, and broadly deployable.

References

- [1] Roman Beltiukov, Wenbo Guo, Arpit Gupta, and Walter Willinger. In search of netunicorn: A data-collection platform to develop generalizable ml models for network security problems. In *Proceedings of the 2023 ACM SIGSAC Conference on Computer and Communications Security*, pages 2217–2231, 2023.

¹The total token cost is based on both input and output tokens, each typically charged at a different rate. For simplicity, we represent it here as the total number of tokens.

- [2] Bluesky. <https://bsky.app>, 2025. Accessed: 2025-08-29.
- [3] Bright Data - All in One Platform for Proxies and Web Scraping. <https://brightdata.com/>.
- [4] Francesco Bronzino, Paul Schmitt, Sara Ayoubi, Guilherme Martins, Renata Teixeira, and Nick Feamster. Inferring streaming video quality from encrypted traffic: Practical models and deployment experience. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 3(3):1–25, 2019.
- [5] Thibault Le Sellier De Chezelles, Maxime Gasse, Alexandre Drouin, Massimo Caccia, Léo Boisvert, Megh Thakkar, Tom Marty, Rim Assouel, Sahar Omidi Shayegan, Lawrence Keunho Jang, Xing Han Lù, Ori Yoran, Dehan Kong, Frank F. Xu, Siva Reddy, Quentin Cappart, Graham Neubig, Ruslan Salakhutdinov, Nicolas Chapados, and Alexandre Lacoste. The browsergym ecosystem for web agent research, 2025.
- [6] Cisco Systems, Inc. Webex. <https://www.webex.com>, 2025. Accessed: 2025-08-29.
- [7] Jaber Daneshamooz, Jessica Nguyen, William Chen, Sanjay Chandrasekaran, Satyandra Guthula, Ankit Gupta, Arpit Gupta, and Walter Willinger. Addressing the ml domain adaptation problem for networking: Realistic and controllable training data generation with netreplica. *arXiv preprint arXiv:2507.13476*, 2025.
- [8] Xiang Deng, Yu Gu, Boyuan Zheng, Shijie Chen, Samuel Stevens, Boshi Wang, Huan Sun, and Yu Su. Mind2web: Towards a generalist agent for the web, 2023.
- [9] Lutfi Eren Erdogan, Nicholas Lee, Sehoon Kim, Suhong Moon, Hiroki Furuta, Gopala Anumanchipalli, Kurt Keutzer, and Amir Gholami. Plan-and-Act: Improving Planning of Agents for Long-Horizon Tasks, April 2025.
- [10] Espn. <https://www.espn.com>, 2025. Accessed: 2025-08-29.
- [11] Gemini 2.5: Our newest gemini model with thinking. <https://blog.google/technology/google-deepmind/gemini-model-thinking-updates-march-2025/>. Accessed: 2025-08-12.
- [12] Google LLC. Google meet. <https://workspace.google.com/resources/video-conferencing/>, 2025. Accessed: 2025-08-29.
- [13] Google LLC. Youtube. <https://www.youtube.com>, 2025. Accessed: 2025-08-29.
- [14] Hulu. <https://www.hulu.com>, 2025. Accessed: 2025-08-29.
- [15] Jitsi. <https://jitsi.org>, 2025. Accessed: 2025-08-29.
- [16] Jing Yu Koh, Robert Lo, Lawrence Jang, Vikram Duvvur, Ming Chong Lim, Po-Yu Huang, Graham Neubig, Shuyan Zhou, Ruslan Salakhutdinov, and Daniel Fried. Visualwebarena: Evaluating multimodal agents on realistic visual web tasks, 2024.
- [17] Langchain. <https://www.langchain.com>. Accessed: 2025-08-23.
- [18] LinkedIn Corporation. LinkedIn. <https://www.linkedin.com>, 2025. Accessed: 2025-08-29.
- [19] Qingzhou Luo, Farah Hariri, Lamyaa Eloussi, and Darko Marinov. An empirical analysis of flaky tests. In *Proceedings of the 22nd ACM SIGSOFT international symposium on foundations of software engineering*, pages 643–653, 2014.
- [20] Meta Platforms, Inc. Instagram. <https://www.instagram.com>, 2025. Accessed: 2025-08-29.
- [21] Microsoft Corporation. Microsoft teams. <https://www.microsoft.com/en-us/microsoft-teams/group-chat-software#Products-and-services>, 2025. Accessed: 2025-08-29.
- [22] Michael Mintz and SeleniumBase contributors. Seleniumbase.

- [23] Magnus Müller and Gregor Žunič. Browser use: Enable ai to control your browser, 2024.
- [24] Udit Paul, Vinothini Gunasekaran, Jiamo Liu, Tejas Narechania, Arpit Gupta, and Elizabeth Belding. Decoding the divide: Analyzing disparities in broadband plans offered by major us isps. In *Proceedings of the ACM SIGCOMM Conference*, SIGCOMM '23, pages 578–591, New York, United States, 2023. Association for Computing Machinery.
- [25] Michael O Rabin and Dana Scott. Finite automata and their decision problems. *IBM journal of research and development*, 3(2):114–125, 1959.
- [26] Reddit, Inc. Reddit. <https://www.reddit.com>, 2025. Accessed: 2025-08-29.
- [27] Ringcentral. <https://www.ringcentral.com>, 2025. Accessed: 2025-08-29.
- [28] Roku, Inc. Roku. <https://www.roku.com>, 2025. Accessed: 2025-08-29.
- [29] Sai Adarsh S, Prasanna Venkateshan, and Prithvi Alva Suresh. Emulating human-like mouse movement using bézier curves and behavioural models for advanced web automation. *IJIRT*, 12(3):1364–1370, 2025.
- [30] Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023.
- [31] Stanford University. Puffer, 2025. Accessed: 2025-08-29.
- [32] Al Sweigart. Pyautogui.
- [33] Talky. <https://about.talky.io>, 2025. Accessed: 2025-08-29.
- [34] The Walt Disney Company. Disney. <https://www.disneyplus.com>, 2025. Accessed: 2025-08-29.
- [35] Swapna Thorve, Chandani Sreshtha, and Na Meng. An empirical study of flaky tests in android apps. In *2018 IEEE International Conference on Software Maintenance and Evolution (ICSME)*, pages 534–538. IEEE, 2018.
- [36] Tubi, Inc. Tubi. <https://tubitv.com>, 2025. Accessed: 2025-08-29.
- [37] Twitch. <https://www.twitch.tv>, 2025. Accessed: 2025-08-29.
- [38] ultrafunkamsterdam. undetected-chromedriver.
- [39] Whereby. <https://whereby.com>, 2025. Accessed: 2025-08-29.
- [40] X. <https://x.com>, 2025. Accessed: 2025-08-29.
- [41] Jianwei Yang, Hao Zhang, Feng Li, Xueyan Zou, Chunyuan Li, and Jianfeng Gao. Set-of-Mark Prompting Unleashes Extraordinary Visual Grounding in GPT-4V, November 2023.
- [42] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models, 2023.
- [43] Shuyan Zhou, Frank F. Xu, Hao Zhu, Xuhui Zhou, Robert Lo, Abishek Sridhar, Xianyi Cheng, Tianyue Ou, Yonatan Bisk, Daniel Fried, Uri Alon, and Graham Neubig. Webarena: A realistic web environment for building autonomous agents, 2024.
- [44] Zoho meeting. <https://www.zoho.com/meeting/>, 2025. Accessed: 2025-08-29.
- [45] Zoom Video Communications, Inc. Zoom. <https://www.zoom.com/?lang=en-US>, 2025. Accessed: 2025-08-29.

A Appendix / supplemental material



Figure 2: ESPN workflow: log into Disney+, select the account, enter the PIN if required, navigate to ESPN, play the first video, and advance the playback slider to the five-minute mark.

Table 1: Evaluation of *NetGent* across different workflows, performed entirely without using cached states.

Application Type	Platform	Tokens ($\times 10^3$)	Price (\$)	Time (s)
Video Conferencing	Zoom [45]	146.8	0.051	138
Video Conferencing	Microsoft Teams [21]	242.9	0.085	238
Video Conferencing	Google Meet [12]	169.7	0.060	160
Video Conferencing	Zoho Meeting [44]	137.6	0.049	190
Video Conferencing	Jitsi [15]	98.5	0.035	120
Video Conferencing	Whereby [39]	77.1	0.026	78
Video Conferencing	Ring Central [27]	131.9	0.047	139
Video Conferencing	Talky [33]	110.9	0.041	129
Video Conferencing	Webex [6]	107.0	0.037	114
Video on Demand + Live stream	YouTube [13]	131.7	0.047	105
Video on Demand + Live stream	ESPN [10]	278.7	0.098	339
Video on Demand + Live stream	ESPN (PIN required)	375.4	0.105	406
Video on Demand	Disney Plus [34]	249.2	0.088	335
Video on Demand	Hulu [14]	245.0	0.089	386
Video on Demand	Roku [28]	162.2	0.059	187
Video on Demand	Tubi [36]	140.6	0.049	158
Live stream	Twitch [37]	134.6	0.044	90
Live stream	Puffer [31]	109.3	0.038	126
Social Media	X (Twitter) [40]	201.7	0.068	158
Social Media	Instagram [20]	199.1	0.069	180
Social Media	LinkedIn [18]	143.8	0.049	138
Social Media	Reddit [26]	739.5	0.225	82
Social Media	Bluesky [2]	98.1	0.033	79
Web Scraping	BQT [24] (30 ISPs)	145*	0.050*	120*

* Approximate average numbers across 30 ISPs.

A.1 Language and API Integration

LangChain [17] provides core capabilities for model input/output, tool invocation, and message abstraction, standardizing interactions with Gemini models via `SystemMessage` and `HumanMessage`. LangGraph implements a control-flow layer, composing agent behaviors into state-machine workflows, and is primarily utilized in our Browser Agent, which includes environment observation, multi-step planning, and code generation and execution. VertexAI enables access to LLMs on Google Cloud, allowing the use of Gemini 2.5 chat models [11] through `langchain-google-vertexai` wrapper.

To enhance the Web Agent’s capabilities and efficiency, we incorporated several key open-source contributions. Notably, we integrated the DOM Marker code from Browser Use [23] to implement Set-of-Mark (SoM)

prompting [41], enabling precise interaction with web elements via visual markers. We also adopted the Planner, Replanner, and Executor prompts from Plan-and-Act [9], allowing the agent to decompose complex tasks into structured planning phases followed by systematic execution.

These integrations provide a robust foundation that leverages proven open-source techniques while supporting the unique requirements of *NetGent*, facilitating reliable and scalable automation of complex web workflows.

A.2 Realism

To support realism, we adapted different methods in the web agent to perform tasks in a way that simulate human-like interactions with the web and avoid bot detection. To mitigate automated behavior detection, our system integrates three layers of anti-bot techniques: **browser stealth**, **movement realism**, and **network control**. At the browser layer, we employ SeleniumBase [22] with undetected-chromedriver [38] to suppress common automation fingerprints. Repeated logins during automated workflows can raise suspicion on services (e.g., Disney+), potentially leading to account freezes or forced two-step verification via email. To mitigate this, we support persistent user profiles via the *user-data-dir* option, ensuring continuity without repeated authentication challenges. These design choices enable our agent to more closely emulate a human-operated browser session, preserving cookies and local storage across repeated tasks. To model user behavior, we introduce human-like movement primitives. Our framework calculates absolute screen coordinates and issues operating system-level mouse and keyboard events using PyAutoGUI [32]. Mouse movement follows Bezier-curve trajectories [29] to simulate natural cursor dynamics. Keystrokes are generated with variable inter-key intervals and pauses. Scrolling and hovering operations are tied to the viewport, preventing the “teleporting” behavior typical of automated agents.

Finally, we utilize a pool of IP addresses provided by Bright Data [3] to distribute requests across different locations and IP addresses. This part is essential for many web scraping applications. Without this mechanism, repeated interactions would appear to originate from a single IP address, potentially resembling a denial-of-service pattern and triggering blocking. By varying network origin, the system avoids triggering such defenses while preserving the appearance of genuine user traffic. Together, these measures allow our system to present itself as both behaviorally indistinguishable from a human operator, which is critical for evaluating the robustness of our workflow synthesis framework under realistic conditions.